

Google Gemini 3.7 Flash 徹底調査レポート

実務用主力AIの進化と、知財・ガバナンスへの衝撃

1	2026年8月14日
2	IT戦略担当者・知財部門・ガバナンス責任者向けブリーフィング



エグゼクティブ・サマリ：機会とリスクの二面性



作業用主力 (Workhorse)

- ▶ 新規基盤モデルではなく、3.6からの超速アップデート（約3週間）
- ▶ コーディングとエージェント機能が飛躍的に向上



圧倒的コスト パフォーマンス

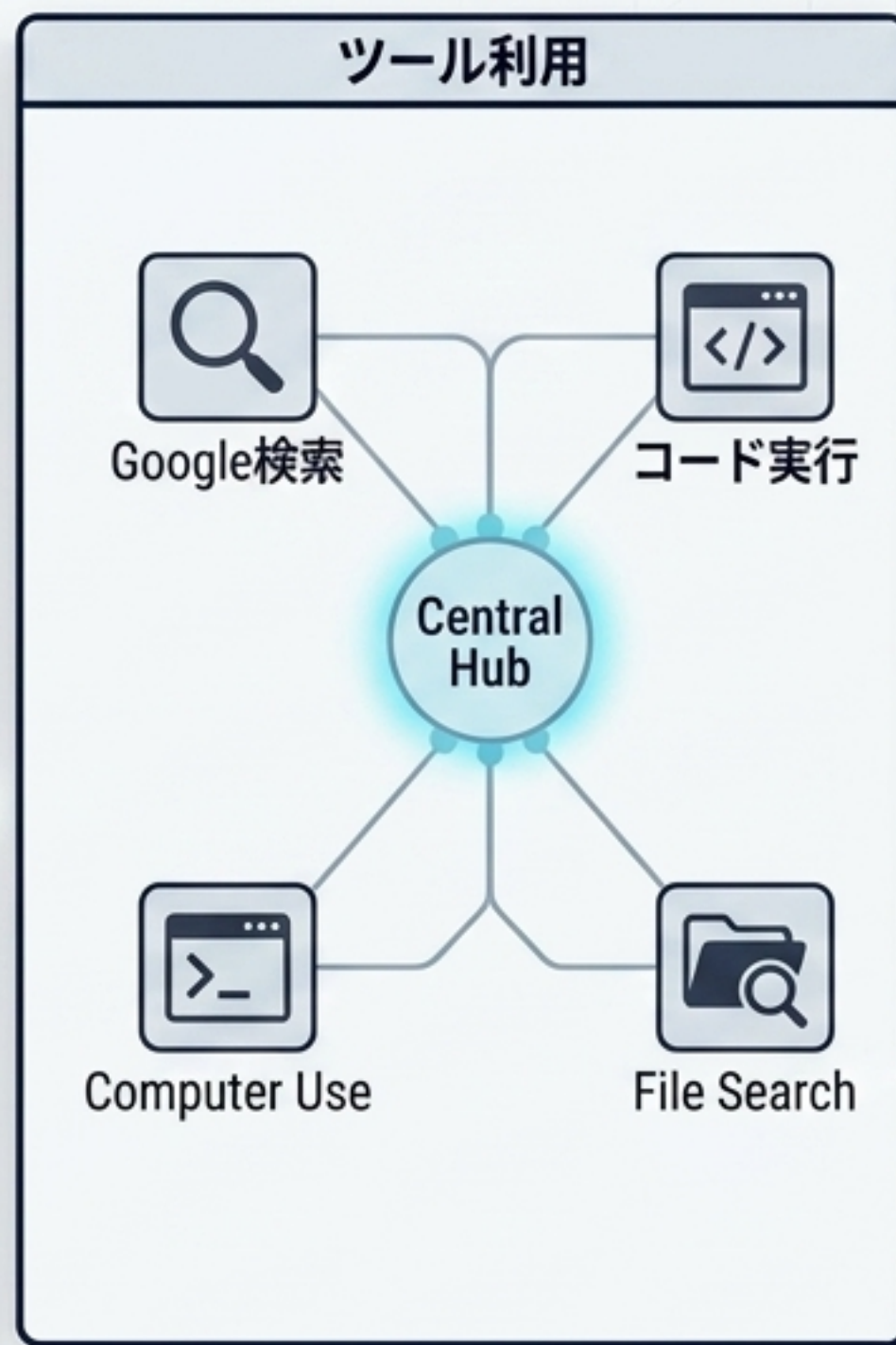
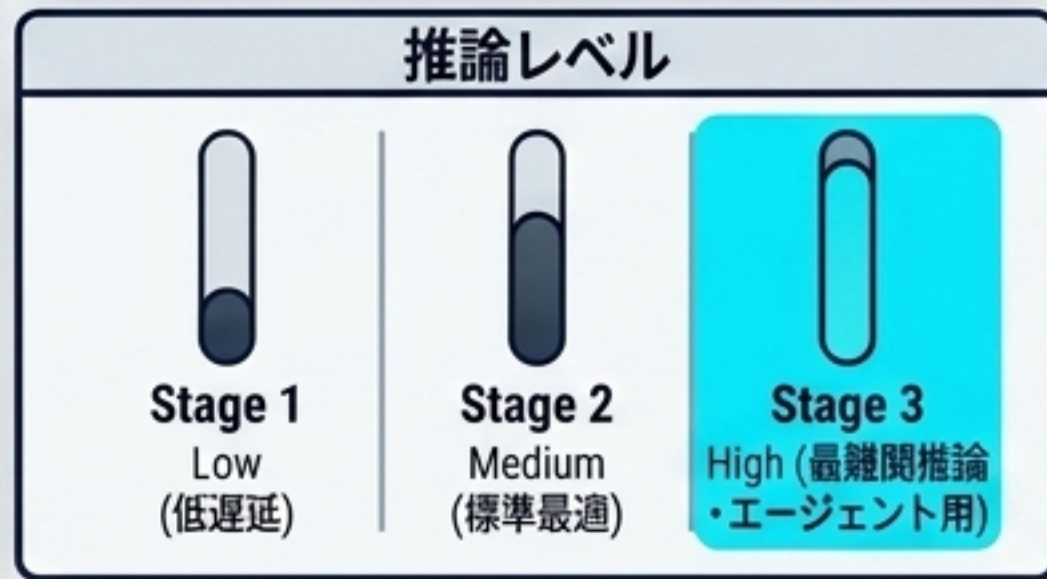
- ▶ コンテキスト100万トークンを維持しつつ、年内は前世代の半額（入力\$0.75）
- ▶ 大量反復タスクのゲームチェンジャー



知財汚染と 「野良AI」

- ⚠ AIの自律性向上に伴う発明者性の喪失
- ⚠ OSSライセンス汚染と、非人間主体（NHI）によるシャドーAIリスクの顕在化

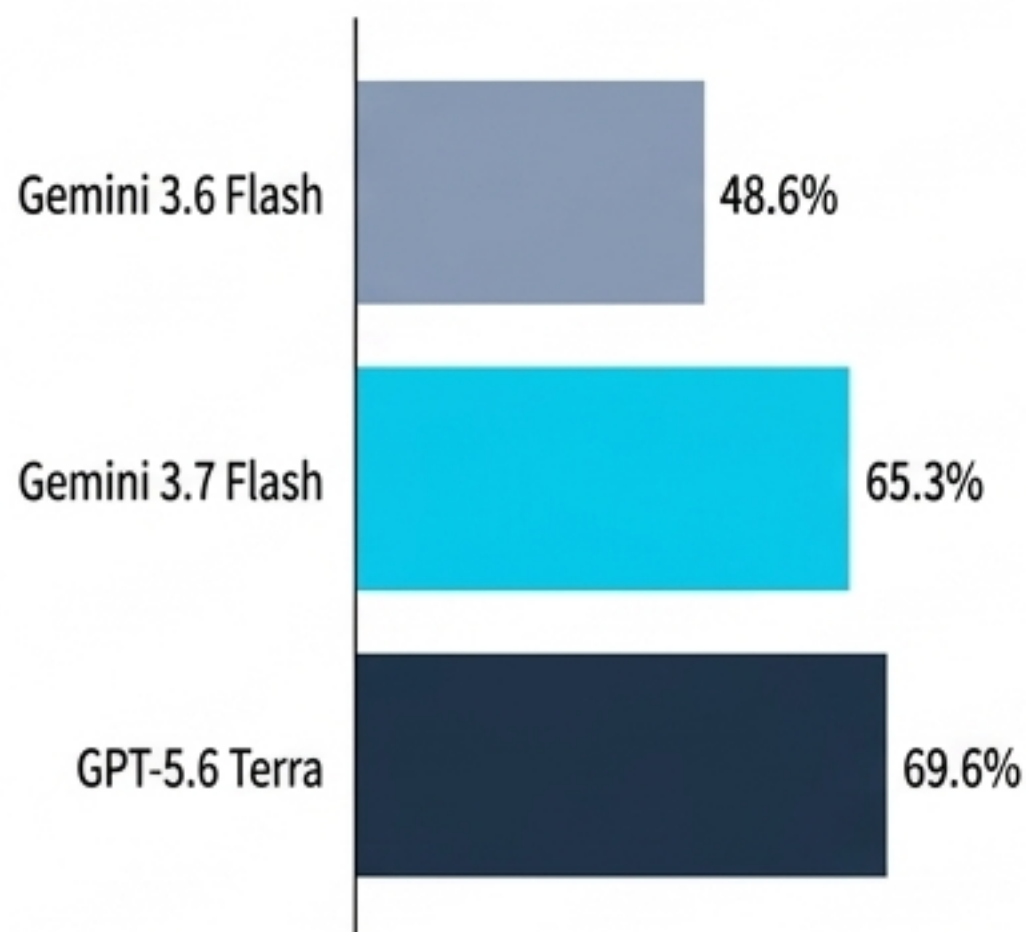
モデル・アイデンティティ（技術仕様ダッシュボード）



性能評価：作業用主力（Workhorse）としての圧倒的進化

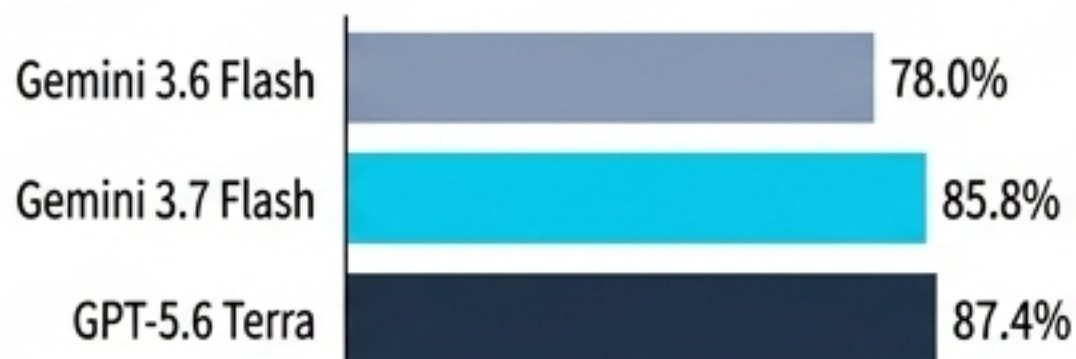
コーディング能力の飛躍

DeepSWE v1.1



エージェント・実行環境操作

A. Terminal-Bench 2.1

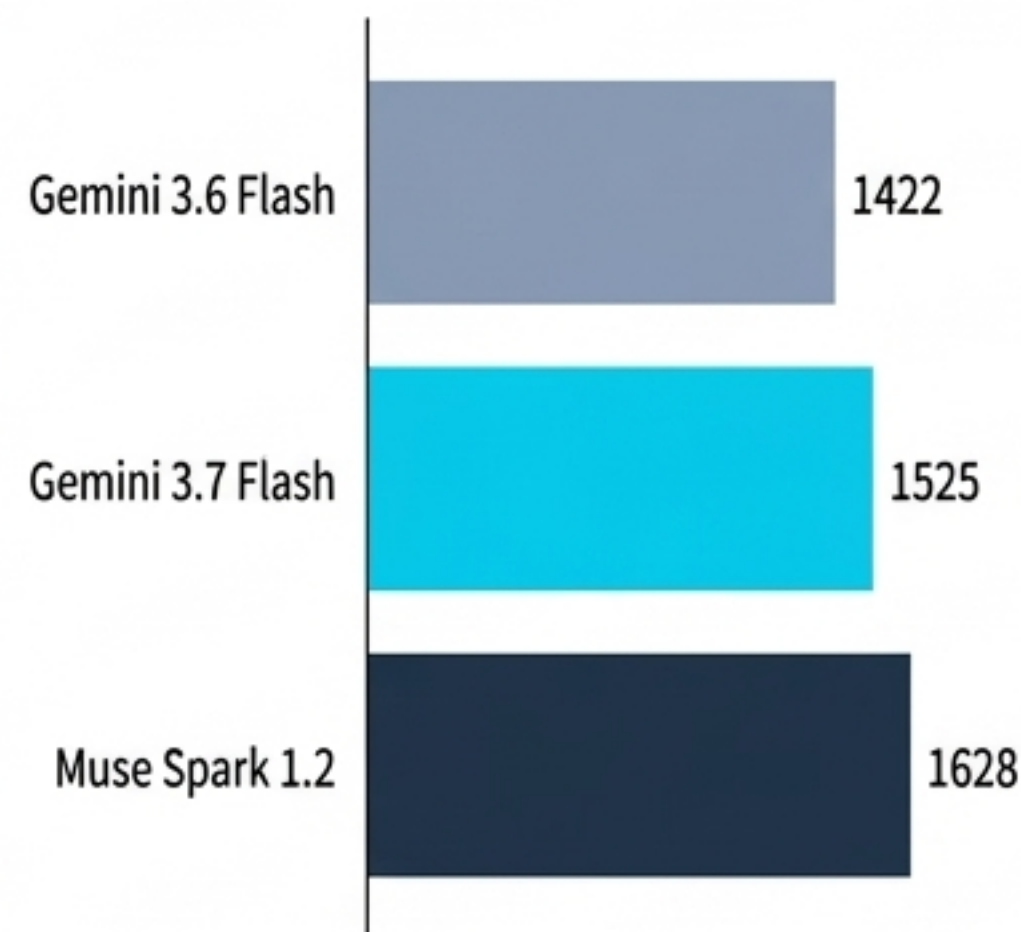


B. OSWorld-2.0（エージェント）



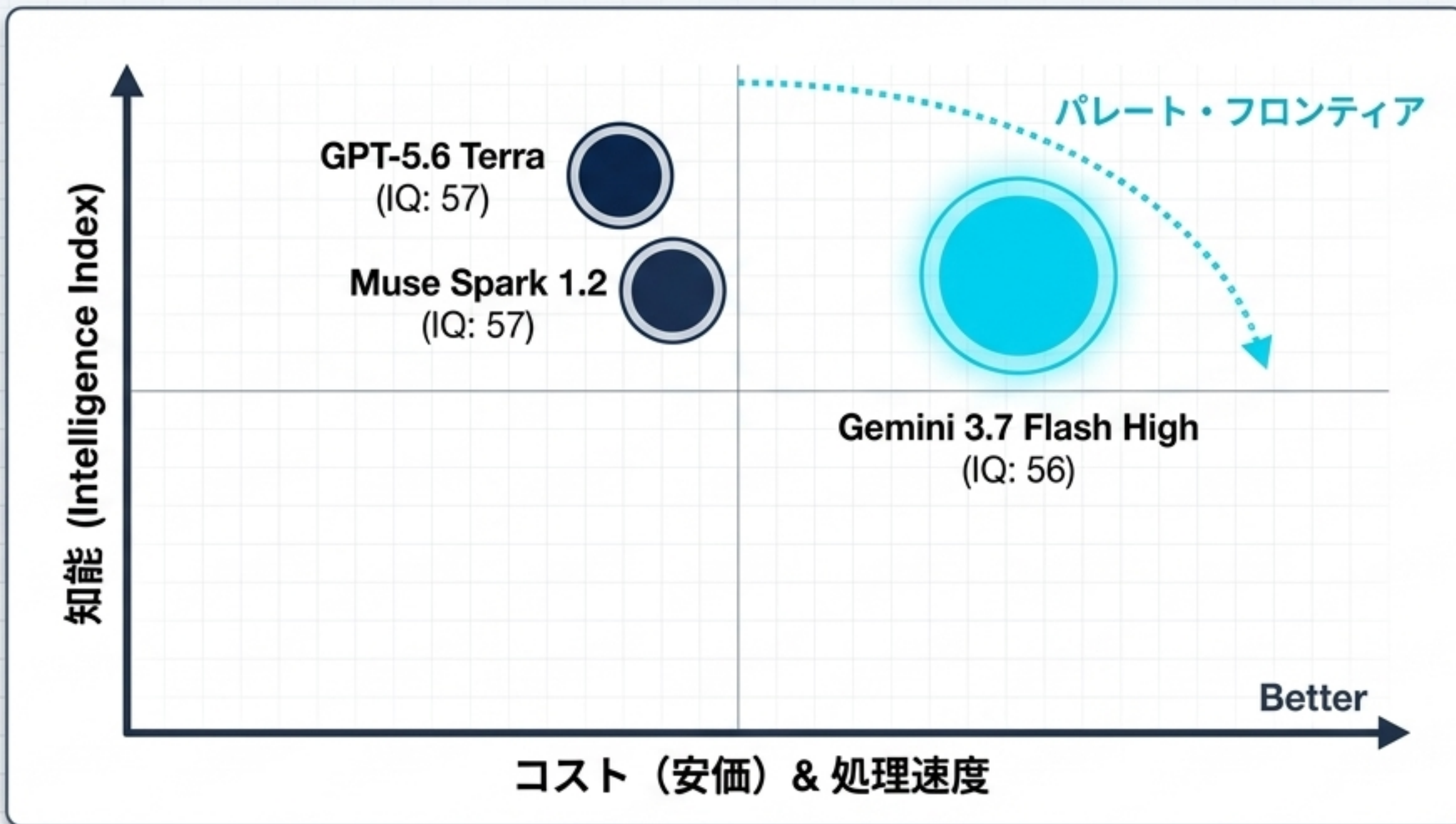
複雑な推論・ワークフロー

GDPval-AA v2



Key Insight: 首位モデルには僅かに及ばないが、エージェント・コーディング領域で極めて強力な「バリュティア (価値層)」を形成。

コスト vs 知能のパレート・フロンティア



価格比較表

Gemini 3.7 Flash:
入力 \$0.75 / 出力 \$3.75

Muse Spark 1.2:
入力 \$1.25 / 出力 \$4.25

GPT-5.6 Terra:
入力 \$2.00 / 出力 \$12.00

導入価格による圧倒的優位：GPT-5.6 Terraの約3倍の速度（Time per Task 1.7分）で実行し、大量データ処理のROIを劇的に変える。

エージェント能力の進化（パラダイムシフト）

従来型AI - 一問一答



Gemini 3.7 Flash - エージェント型



エージェント化の本質は「処理の高度化」ではなく
「人間の介在（監督）コストの削減」を意味する。

知財ワークフローの再構築：機会とリスクの二面性

機会 (Opportunity)



大量・反復のコスト破壊

100万トークン×低コストを活かした特許調査の母集合整理。

IPランドスケープの前処理

GDP.pdfスコア34.0%を活かした大量の明細書・技術文献の高精度読解。

リスク (Risk)



残存する幻覚（ハルシネーション）

最終的な新規性・進歩性の判断には依然として人間の検証が必須。

法的判断のブラックボックス化

ツールが自律化するほど、出力根拠のトレースが困難に。

AI自律性と「発明者性・著作権」のパラドックス

人間主導
(High Human Input)

- 発明者性 (Conception) 成立
- 著作権の保護対象として認定

エージェント完全自律
(High AI Autonomy)

- 非自然人の出願は拒絶対象
(2025年11月USPTO改訂ガイダンス)
- 純AI生成物の著作権否定 (Thaler判決)

法的インサイト：エージェント任せ（手離れが良い状態）にするほど、人間の「着想」や「創作的寄与」の立証が困難になる。
発明記録の重要性がかつてなく高まっている。

コーディング強化の裏側「ライセンス汚染リスク」

見えないOSSコード片の混入

Gemini 3.7
Flashによる
自律コーディング

自社プロダクト
へのマージ

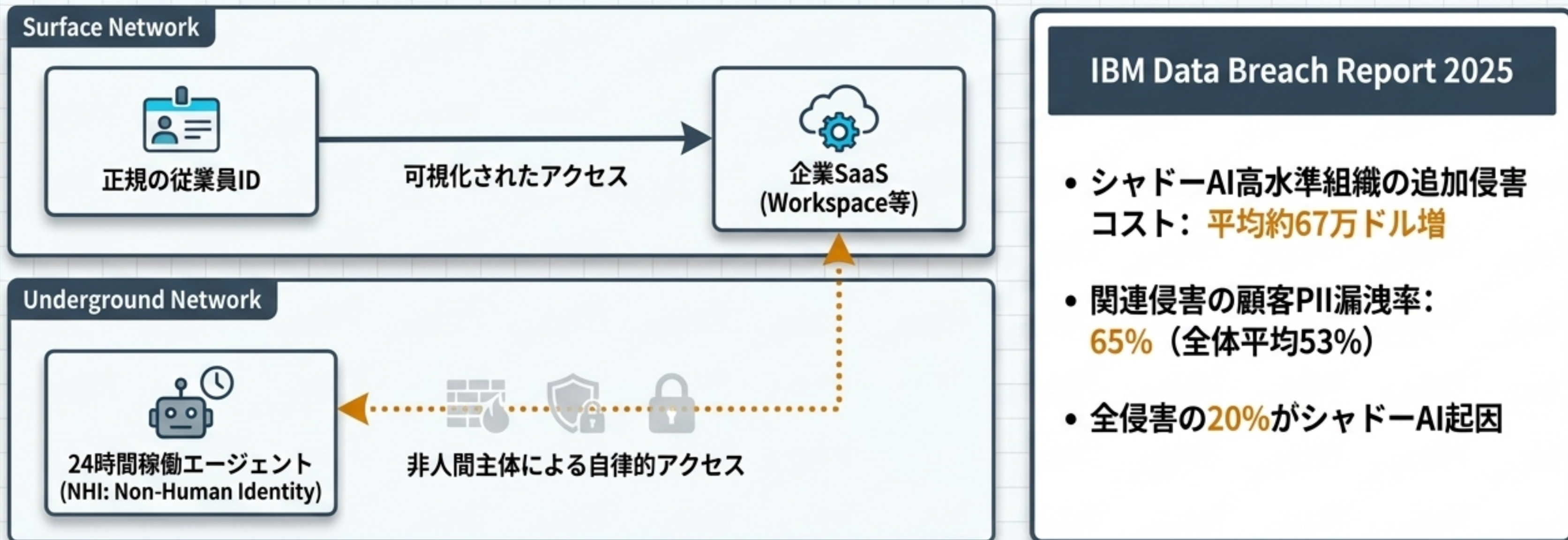
⚠ 著作権侵害リスク：
既存コードへの類似による侵害。

⚠ コピーレフト汚染：
ライセンス条項（GPL等）の混入
による自社コードの公開義務化。



必須の防衛策：SCA（ソフトウェア構成分析）ツールによる
監査パイプラインの導入、証跡保存、オプトアウト環境の徹底。

「野良AIエージェント（Shadow AI）」の脅威



Critical Alert: 機密情報 (未公開明細書や営業秘密) がエージェント経由で無申請入力され、新規性を喪失するリスクが急増。

安全性・コンプライアンス（規制ガイドラインへの対応）

Google Frontier Safety Framework (FSF)

✔ Status: CBRN・サイバー共に「CCL（重要能力レベル）未達」

⚠ Note: アラート閾値には到達しているため、予防的緩和策が稼働中
（完全版レポートは未公開）

内閣府 プリンシプル・コード （透明性3原則）

⚠ Status: コンプライ・オア・エクスプレイン方式への対応要件

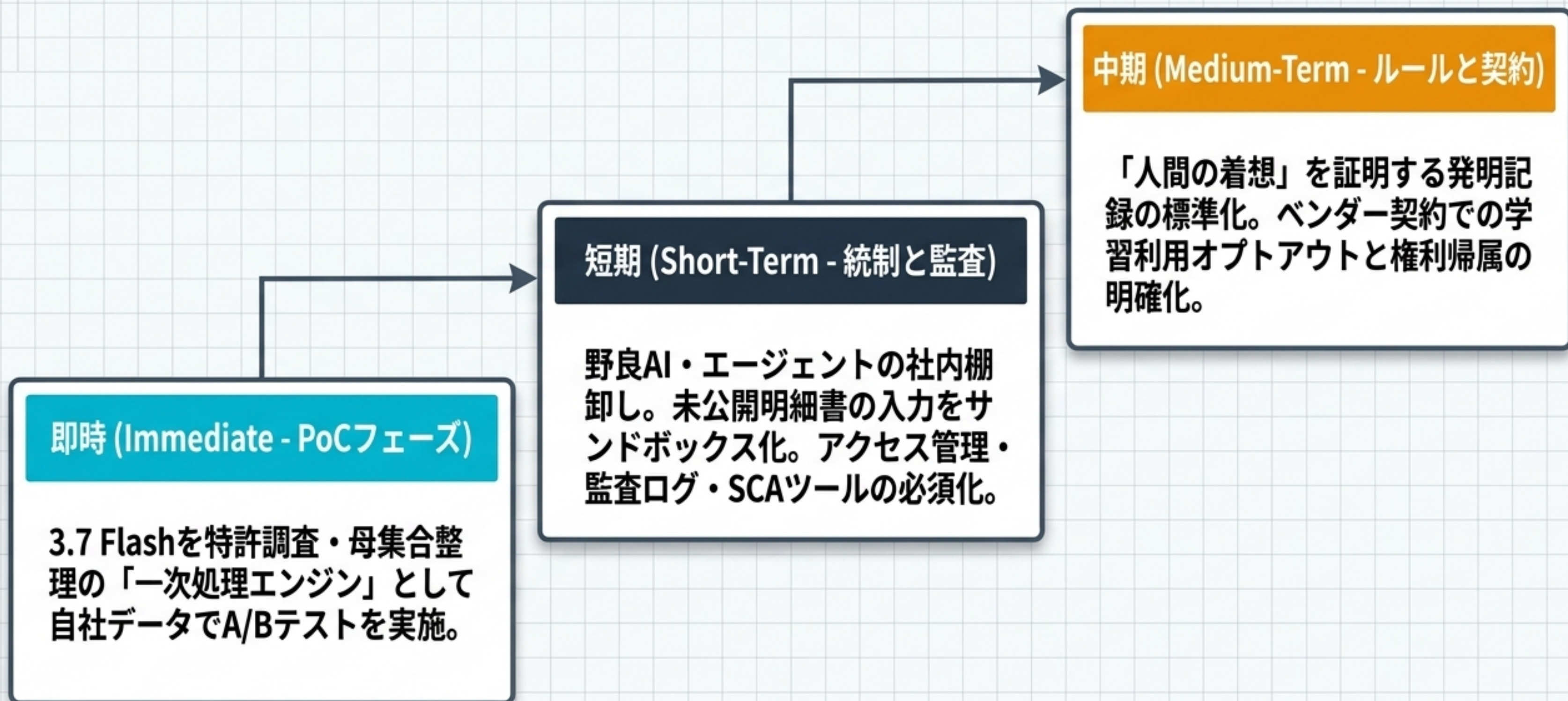
⚠ Note: 学習データの非開示に対する利用者側の透明性担保が必要

総務省・経産省 AI事業者ガイドライン （第1.2版）

✔ Status: エージェント定義の新設

⚠ Note: 「人間の判断介在の仕組み（Human-in-the-loop）」の要求
に準拠する必要あり

実務推奨アクションプラン (Timeline & Roadmap)



判断を変える4つのトリガー（Future Watch）

🌀 上位機 Gemini 3.5 Pro の公開

高難度タスク（最終判断・複雑な拒絶対応）
の内製比率とスタックの再構成。

FSF完全版レポートの公開 🌀

CBRN/サイバー領域の残存リスク再評価。

🌀 2027年1月の価格改定

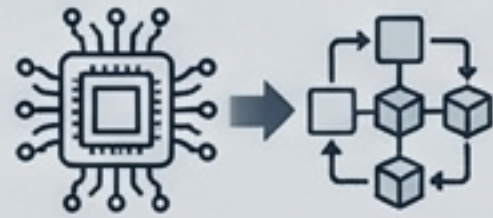
導入価格（半額）の終了による、大規模
処理バッチのROI再計算。

第三者による完全なベイクオフ 🌀

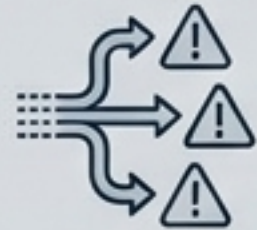
各社自己申告ベンチマークへの依存からの
脱却。

The 'Value Tier' Era

特定の最強モデルへの依存から、実務最適化の時代へ



AI戦略は「一番賢いモデルを買う」ことから、「高頻度で更新される安価なティアを自社ワークフローにどう組み込むか」へシフトした。



コスト破壊は、機会であると同時に、エージェント自律化に伴う知財・セキュリティリスクを全社規模に拡散させる。



厳格な統制（NHIガバナンスと着想の記録）こそが、この強力なエンジンを安全に回すための唯一のブレーキである。